

## HOW UNIVERSITY STAFF EVALUATE GENERATIVE AI: COGNITIVE AND ETHICAL PERSPECTIVES ON TEACHING, TRUST, AND ACADEMIC INTEGRITY

Serhiy Lyeonov

Silesian University of Technology, Poland;  
Lithuania Business College, Lithuania  
ORCID 0000-0001-5639-3008

Artem Artyukhov

Bratislava University of Economics and  
Business, Slovak Republic;  
Sumy State University, Ukraine  
ORCID 0000-0003-1112-6891

Svitlana Bilan

Rzeszów University of Technology, Poland;  
Széchenyi István University, Győr, Hungary  
ORCID 0000-0001-9814-5459

Zoltán Bács

University of Debrecen  
Hungary  
ORCID 0000-0003-0612-658X

**Abstract:** *The rapid diffusion of generative artificial intelligence (GenAI) is reshaping higher education by challenging traditional roles of teaching, trust, and academic integrity. This study aims to explore how university staff cognitively and ethically evaluate GenAI by analysing perceptions of its pedagogical replacement potential, practical feasibility, academic integrity risks, and perceived reliability across national contexts. The analysis is based on an anonymous cross-sectional survey of 637 respondents conducted between May and September 2025, using descriptive statistics, correlations, regression models, and exploratory factor analysis. The findings show that perceived replacement potential is low ( $M = 2.47$ ), with over 51% of respondents rejecting the idea of AI replacing teachers. Academic integrity concerns are the strongest dimension ( $M = 3.51$ ), while trust in AI accuracy remains low ( $M = 1.99$ ), indicating widespread scepticism. Perceived cost and complexity do not significantly influence beliefs about replacement ( $R^2 = 0.007$ ;  $p = 0.117$ ), suggesting a weak relationship between feasibility and perceived impact. Finally, moderate positive correlation ( $\rho = 0.34$ ) and low reliability ( $\alpha = 0.50$ ;  $\alpha = 0.45$ ) provide evidence that perceptions of GenAI are fragmented and multidimensional rather than internally consistent.*

**Keywords:** *generative artificial intelligence, higher education, academic integrity, AI trust, technology adoption, cognitive perceptions, digital transformation.*



## INTRODUCTION

The rapid expansion of generative artificial intelligence (GenAI) is transforming higher education systems worldwide, reshaping how knowledge is produced, delivered, and assessed. International organisations emphasise that AI is becoming a central pillar of digital transformation in education, influencing teaching practices, learning environments, and institutional governance. In particular, the European Commission highlights that AI-driven tools accelerate personalised, data-intensive education while redefining educators' roles (European Commission, 2026). Similarly, the World Bank stresses that digital technologies, including AI, support resilient and adaptive education systems (Rajasekaran et al., 2024).

Importantly, the integration of GenAI is not only a technological process but a human-centred transformation. From this perspective, the impact of AI depends on how it is perceived and evaluated by individuals within institutional contexts. University staff act as both users and gatekeepers of AI, shaping its implementation and legitimacy. Their perceptions, therefore, represent a key mechanism through which technological potential translates into real educational outcomes. Within this framework, two analytically distinct dimensions are central: cognitive and ethical evaluation. Cognitive evaluation refers to perceived capability, usefulness, and reliability of GenAI (e.g., accuracy and replacement potential), while ethical evaluation concerns academic integrity, responsibility, and risks of misuse. These dimensions may not be significantly associated, suggesting that attitudes towards GenAI are multidimensional rather than internally consistent.

At the same time, GenAI raises fundamental questions about trust, reliability, and transparency. While AI offers significant benefits, it also introduces risks such as misinformation, opacity, and algorithmic bias (OECD, n.d.). These concerns are particularly salient in higher education, where knowledge credibility is essential. Accordingly, UNESCO emphasises that AI use must be grounded in human oversight, accountability, and trustworthiness (UNESCO, 2022). Ethical challenges are especially pronounced, as GenAI blurs boundaries of authorship and originality, creating new forms of academic misconduct (UNESCO, 2023; European Commission, 2026).

Despite strong policy attention, a key gap remains in understanding how university staff cognitively and ethically evaluate GenAI in practice. Existing frameworks provide normative guidance but limited empirical insight into how different perception dimensions relate to each other. Moreover, attitudes towards AI are often treated as unified, without testing their internal structure. To address this gap, this study operationalises GenAI perceptions through five measurable dimensions: perceived replacement potential, cost, complexity, academic integrity risks, and perceived accuracy. This approach enables a structured and empirically grounded assessment of both cognitive and ethical evaluations, as well as their internal consistency. Therefore, examining how university staff perceive GenAI across different national and institutional contexts is essential for bridging the gap between policy frameworks and real-world implementation, highlighting that digital transformation depends not only on technology itself but on how it is cognitively interpreted and ethically assessed (OECD, n.d.; Rajasekaran et al., 2024).

## LITERATURE REVIEW

The growing integration of generative artificial intelligence (GenAI) into higher education reflects four core dimensions directly relevant to perceptions of GenAI, organisational practices, and human–technology interaction. The existing scientific landscape demonstrates that AI is not merely a technological innovation but a systemic driver of change affecting teaching roles, trust dynamics, ethical governance, and institutional adoption processes (Bilan et al., 2022; Koibichuk et al., 2023; Škare et al., 2025; Zimosz & Ober, 2025). Within this context, the evaluation of GenAI by university staff emerges as a multidimensional phenomenon shaped by cognitive, ethical, and organisational factors.

### **Transformation of the teacher's role: augmentation vs substitution**

A central stream of the literature examines how AI reshapes educators' roles, particularly in terms of augmentation versus substitution. Evidence consistently suggests that AI does not fully replace teachers but rather transforms their functions, shifting them towards facilitation, critical evaluation, and interaction management (Celik et al., 2022; Hwang et al., 2020; Jeon & Lee, 2023; Zhang & Fan, 2026). This transformation is embedded in broader processes of digitalisation and innovation transfer in education, where AI contributes to new pedagogical models and learning environments (González-Campos et al., 2024; Rahmanov et al., 2025; Skydan, 2024).

At the same time, perceptions of substitution risk remain an important dimension of AI evaluation, particularly as generative systems increasingly demonstrate advanced cognitive capabilities. The literature highlights that staff attitudes towards AI replacement are shaped not only by technological capability but also by pedagogical beliefs, professional identity, and perceived value of human interaction (Derakhshan & Ghiasvand, 2024; Kohnke & Zou, 2025; Kundu et al., 2025). Moreover, the development of AI-related competencies and labour-market adaptation further reinforces the view that AI leads to role reconfiguration rather than displacement (Bîrcă et al., 2025; Popa et al., 2024; Zuo et al., 2025).

### **Trust in AI: perceived reliability, accuracy, and human–AI interaction**

A second core stream focuses on trust as a central determinant of AI adoption and evaluation. Trust in AI is conceptualised as a multidimensional construct influenced by perceived reliability, transparency, and risk, as well as by individual experience and contextual factors (Grabowska & Rzemieniak, 2025; Nazaretsky et al., 2022). In educational settings, trust is particularly important because it directly affects educators' willingness to rely on AI-generated outputs and integrate them into teaching practices.

The literature suggests that trust in AI does not emerge automatically but requires calibration between perceived capabilities and perceived risks. While AI systems offer efficiency and decision-support benefits (Kozhakhmetova et al., 2024), concerns related to opacity, errors, and bias may limit trust and adoption (Ejdys et al., 2025; Roba & Moulay, 2024; Qian, 2021). At the same time, human–AI interaction is increasingly understood as a relational process in which trust is co-shaped by user experience, pedagogical context, and emotional factors such as enjoyment and perceived usefulness (Moravec et al., 2024;

Yarovenko et al., 2024; Zhang & Fan, 2026). These findings indicate that perceptions of AI accuracy and reliability are central to understanding how university staff cognitively evaluate GenAI.

### **Academic integrity and ethical governance as sociotechnical challenges**

A third major stream addresses ethical concerns, particularly academic integrity, as a core dimension of AI evaluation. The rapid development of generative AI has introduced new forms of academic misconduct, including AI-assisted plagiarism and authorship ambiguity, challenging existing regulatory frameworks (Cotton et al., 2023; Eke, 2023; Sullivan et al., 2023; Yusuf et al., 2024). These challenges are not purely technical but sociotechnical, requiring coordinated responses involving policy, pedagogy, and institutional governance.

The literature emphasises that academic integrity is shaped by both individual ethical awareness and broader institutional conditions, including governance quality and educational experience (Artyukhov et al., 2024; Borissov & Liuta, 2025). Ethical AI use in education is therefore grounded in principles such as accountability, transparency, and human oversight, which aim to safeguard fairness and trust in knowledge production (Mura & Stehlíková, 2025; Mujtaba, 2025; Nguyen et al., 2022). Furthermore, the integration of AI into research and teaching raises concerns about data ethics, methodological reliability, and responsible innovation (Ejdys et al., 2025; Hutchinson, 2025).

Importantly, ethical perceptions are not necessarily aligned with cognitive evaluations of AI performance. Evidence suggests that individuals may simultaneously recognise the usefulness of AI while expressing concerns about its misuse, indicating that ethical and functional assessments follow different evaluative logics (Ioan-Franc & Diamescu, 2026; Iordache et al., 2025). This supports the need to analyse academic integrity as a distinct dimension of GenAI perception.

### **Organisational adoption, leadership, and perception heterogeneity**

A fourth stream focuses on AI adoption as an organisational and context-dependent process. The successful integration of AI in higher education depends not only on technological readiness but also on leadership capacity, institutional culture, and stakeholder engagement (Pálffy & Mihályka, 2024; Némethová et al., 2025; Welch, 2025). Educational institutions operate as complex systems where perceptions of AI are shaped by organisational structures, governance models, and external environments (Pálffy et al., 2023; Kuzior et al., 2024; Mong & Thanh, 2024).

Empirical evidence highlights substantial heterogeneity in AI perceptions across individuals, groups, and contexts. Differences in digital literacy, awareness, and prior experience contribute to varying attitudes towards AI adoption and use (Kundu et al., 2025; Moravec et al., 2024; Pisičá & Zaharia, 2025). Moreover, the broader socio-economic and sectoral diffusion of AI underscores the need to understand AI adoption as part of a wider transformation across multiple domains, including education, management, and the creative industries (Kaczorowska-Spychalska et al., 2024; Schinello, 2025; Yarovenko et al., 2024).

This perspective aligns with the view that AI perceptions are embedded in human-centred and institutional contexts, where cognitive, ethical, and practical considerations interact but do not necessarily form a unified evaluative framework.

The existing scientific landscape demonstrates that evaluating generative AI in higher education is inherently multidimensional, encompassing perceptions of role transformation, trust, ethical risks, and organisational context. However, prior research has largely examined these dimensions separately, with limited attention to their internal structure and interrelationships. Therefore, there remains a need for integrative empirical studies that simultaneously capture cognitive and ethical evaluations of GenAI among university staff and assess whether these perceptions form a coherent or fragmented attitudinal framework.

This study aims to explore how university staff cognitively and ethically evaluate generative artificial intelligence in higher education by analysing perceptions of its pedagogical replacement potential, practical feasibility, academic integrity risks, and perceived reliability across different national and contextual settings.

## **METHODS**

This study employed a quantitative cross-sectional survey design to examine how university staff perceive the opportunities, barriers, and risks associated with GenAI in higher education. The analytical focus was placed on attitudes towards teacher replacement, implementation feasibility, academic integrity, perceived accuracy, and cross-group variation in these perceptions. Although the broader conceptual framework initially included six research questions, only five were subjected to full empirical testing in the final analysis. One initially formulated research question was not examined as a separate empirical block because, after instrument reduction and final variable selection, the available items did not provide a sufficiently distinct operational basis for robust standalone testing. The reported analyses, therefore, correspond to the five research questions that can be directly operationalised using the final survey variables.

### **Research design and data collection**

Data were collected through an anonymous online questionnaire administered between 12 May 2025 and 10 September 2025. The survey remained open for approximately four months, allowing responses to accumulate across multiple national contexts and institutional environments. The anonymous format was chosen to reduce social desirability bias and to encourage respondents to express their views freely on potentially sensitive professional issues, such as the possible replacement of teachers by AI and concerns about academic integrity. Importantly, the questionnaire did not request any direct personal identifiers, such as names, email addresses, telephone numbers, home addresses, institutional IDs, or other identifying information. Nor did it include sensitive questions concerning health status, religion, political affiliation, ethnicity, sexual orientation, income, or other protected personal characteristics. As a result, the study relied on a non-identifying and low-risk survey design.

The final analytical dataset contained 637 valid observations and seven variables: one country variable, one binary self-identification variable indicating whether the respondent was Ukrainian, and five Likert-type perception items covering GenAI replacement potential, implementation cost, complexity, academic integrity risks, and perceived accuracy. The final reduced instrument reflected the decision to retain only those variables directly relevant to the tested research questions and sufficiently complete for comparative and multivariate analysis.

## Survey instrument and measurement

The final questionnaire comprised six substantive items and one background item (Appendix A). The country of organisational location was included as a contextual variable, while the five analytical perception variables were measured on a 5-point Likert scale ranging from 1 = strongly disagree to 5 = strongly agree. The dependent and explanatory variables were operationalised as follows. Given the ordinal nature of the variables, results are interpreted cautiously, and consistency across parametric and non-parametric models is used as the primary validity criterion. To strengthen measurement validity, multiple complementary analytical approaches were employed. Convergent evidence was assessed through the consistency of results across non-parametric correlations, linear regression models, and ordinal logistic models. Internal consistency of constructs was evaluated using Cronbach's alpha, while exploratory factor analysis was used to examine the latent structure of the perception variables. This triangulation of methods enhances the robustness of the findings and supports the interpretation of GenAI perceptions as multidimensional rather than as a single unified construct.

The variable replacement measured agreement with the statement that GenAI will replace university teachers in the future. This variable was used as the main outcome for analysing perceived threat to academic teaching roles. The variables cost and complexity captured perceived implementation barriers. The variable integrity measured the extent to which respondents believed that GenAI leads to violations of academic integrity. The variable accuracy reflected perceived trust in the correctness and reliability of AI-generated information. The binary variable Ukrainian was used for subgroup comparison, while country served as the basis for cross-national analysis. In the final sample, the Ukrainian-status variable was coded as a two-category factor ("Yes"/"No"), and country was treated as a categorical grouping variable.

## Data preparation

Before analysis, the dataset was cleaned and simplified. Only the variables retained in the final reduced instrument were preserved. The Likert-type items were converted into numeric form to support descriptive, correlational, and regression-based procedures. The Ukrainian-status variable was converted into a factor for group comparisons, while the country variable remained categorical. For country-level analysis, observations were further restricted to countries with at least 10 cases to reduce instability in group-level estimates and avoid drawing conclusions from very small national subsamples. This threshold produced a comparative country subsample of 19 country categories in the country-level analysis.

Because the survey relied on self-reported perceptions measured on five-point ordinal scales, particular attention was paid to the choice of analytical techniques. Descriptive statistics were first examined to assess means, standard deviations, medians, skewness, and kurtosis. Although the variables did not exhibit extreme distributional irregularities, the ordinal nature of the data remained central to the analytical strategy. For this reason, non-parametric tests and ordinal models were employed where substantively appropriate, while linear models were retained as complementary approximations for ease of interpretation and comparability.

## Analytical strategy

The empirical analysis was organised around five tested research questions.

For RQ1, which asked to what extent university staff perceive GenAI as a threat to academic teaching roles, the analysis relied on descriptive statistics and response distributions for the replacement variable. This included the mean, standard deviation, median, and the percentage distribution across the five response categories. The purpose was to determine whether perceptions were concentrated in disagreement, neutrality, or agreement.

For RQ2, which examined whether perceptions of cost and complexity influence beliefs about GenAI replacing teachers, bivariate Spearman rank-order correlations were first calculated between replacement and each implementation barrier variable. These correlations were followed by a multiple linear regression model with replacement as the dependent variable and cost and complexity as explanatory variables. To account for the ordinal character of the dependent variable, an ordinal logistic regression model was also estimated as a robustness check. This combined strategy enabled assessment of both simple associations and the joint predictive contribution of perceived implementation barriers.

For RQ3, which addressed whether perceptions of academic integrity violations are associated with distrust in GenAI accuracy, the analysis again began with a Spearman rank correlation between integrity and accuracy. A simple linear regression with accuracy as the dependent variable and integrity as the predictor, followed by an ordinal logistic regression to test whether the null association persisted under an ordinal specification. An extended linear model, adding cost and complexity as controls, was also estimated to assess whether the relationship between integrity and perceived accuracy remained absent when other practical considerations were included. OLS estimates are reported for interpretability, while ordinal models are treated as the primary specification for inference.

For RQ4, which concerned differences across countries and between Ukrainian and non-Ukrainian respondents, two levels of analysis were performed. First, descriptive group means and Welch two-sample t-tests were used to compare Ukrainian and non-Ukrainian respondents across all five GenAI perception variables. These models were supplemented with simple linear regressions, with Ukrainian status as the sole predictor, to explicitly express directional contrasts. Second, for the country-level analysis, descriptive statistics were calculated for countries with at least 10 respondents. Kruskal–Wallis rank-sum tests were then applied to test whether distributions differed significantly across countries. Where statistically significant overall country differences were detected, post hoc Dunn tests with Bonferroni adjustment were used to identify specific pairwise contrasts. In addition, a linear model with country dummy variables was estimated for the replacement variable, with Austria as the reference category, to provide interpretable country contrasts.

For RQ5, which examined whether perceptions of GenAI reflect a unified attitudinal construct or a multidimensional framework, a correlation-based and psychometric strategy was adopted. Spearman's rank correlations were first calculated across all five perception variables to assess whether conceptually related items moved together. Cronbach's alpha was then computed for two tentative attitudinal blocks: a capability-oriented block (replacement and accuracy) and a sceptical block (cost, complexity, integrity). Because the observed reliability coefficients were low, this part of the analysis primarily assessed whether a unified attitudinal structure could be empirically justified. Finally, exploratory factor analysis was conducted to

examine whether the variables clustered into latent dimensions. Before extraction, Kaiser–Meyer–Olkin and Bartlett’s tests were used to assess factorability. A forced two-factor solution with Oblimin rotation was then estimated to evaluate whether the data approximated separate dimensions of capability and barriers. Heatmap and scatterplot visualisations were additionally used to support the interpretation of the correlation structure and the relationship between integrity and perceived accuracy.

### **Statistical procedures and software**

All analyses were conducted in R. Descriptive statistics were generated using standard summary procedures and the psych package. Correlations were estimated using Spearman’s rho to avoid reliance on strict interval-scale assumptions. Group comparisons were performed using Welch t-tests and non-parametric Kruskal–Wallis tests where appropriate. Regression models were estimated using ordinary least squares, while ordinal logistic models were estimated to account for ordered response categories. Reliability analysis employed Cronbach’s alpha, and exploratory factor analysis relied on maximum-likelihood extraction with oblimin rotation. Dunn’s post hoc comparisons were adjusted using the Bonferroni method to reduce the risk of Type I error in multiple pairwise comparisons.

### **Ethical considerations**

The study used an anonymous, minimal-risk survey design. No direct identifiers were collected, and no sensitive personal data was requested. The questionnaire focused exclusively on professional perceptions of GenAI in higher education, along with broad contextual indicators such as country of organisational location and whether the respondent identified as Ukrainian. These items were included solely for analytical grouping and comparative interpretation. Because the dataset did not contain identifiable or special-category personal information, the study design minimised privacy risks and protected respondents’ confidentiality. Results were reported only in aggregate, and no attempt was made to identify individual participants or institutions.

### **Methodological limitations**

Several limitations should be acknowledged. First, the study relied on self-reported perceptions rather than observed behaviour, so the findings reflect attitudinal evaluations rather than actual GenAI usage practices. Second, the survey was cross-sectional, so the results capture perceptions at a single point in time rather than changes over time. Third, the Ukrainian subgroup was considerably smaller than the non-Ukrainian subgroup, which limits the balance of comparative inference at that level. Fourth, while country-level comparisons revealed significant heterogeneity, some national subsamples remained modest in size even after applying the 10-respondent threshold. Finally, the psychometric analyses indicate that perceptions of GenAI are only weakly integrated into a common latent structure, which limits the use of composite attitudinal indices and instead supports the interpretation of GenAI perception as a fragmented multidimensional construct.



## RESULTS

The descriptive statistics, based on 637 observations, indicate heterogeneous perceptions of generative AI across different dimensions of its use in higher education. The perceived ability of GenAI to replace university teachers remains relatively low ( $M = 2.47$ ,  $SD = 1.12$ ), with the median equal to 2, suggesting that respondents generally do not support the notion of full substitution of academic staff by AI technologies. The slightly positive skewness (0.32) indicates a modest concentration of responses at lower agreement levels.

**Table 1.** Descriptive statistics of key variables. [Source: Created by the authors.]

| Variable    | N   | Mean | SD   | Median | Min | Max | Skew  | Kurtosis |
|-------------|-----|------|------|--------|-----|-----|-------|----------|
| Replacement | 637 | 2.47 | 1.12 | 2      | 1   | 5   | 0.32  | -0.64    |
| Cost        | 637 | 2.80 | 1.03 | 3      | 1   | 5   | 0.25  | -0.28    |
| Complexity  | 637 | 2.80 | 1.09 | 3      | 1   | 5   | 0.32  | -0.59    |
| Integrity   | 637 | 3.51 | 1.02 | 4      | 1   | 5   | -0.36 | -0.29    |
| Accuracy    | 637 | 1.99 | 1.06 | 2      | 1   | 5   | 0.89  | 0.08     |

*Note.* SD – standard deviation. All variables are measured on a five-point Likert scale (1 – strongly disagree; 5 – strongly agree).

Perceived implementation barriers are moderate. Both cost ( $M = 2.80$ ,  $SD = 1.03$ ) and complexity ( $M = 2.80$ ,  $SD = 1.09$ ) have median values of 3, indicating a neutral to moderately positive assessment of these constraints. The relatively low skewness values (0.25 and 0.32, respectively) suggest balanced distributions, while the negative kurtosis indicates flatter distributions than the normal distribution, reflecting diverse respondent opinions.

Concerns related to academic integrity emerge as the most pronounced dimension. The integrity variable exhibits the highest mean ( $M = 3.51$ ,  $SD = 1.02$ ) and a median of 4, indicating that respondents tend to agree that GenAI may lead to violations of academic integrity. The negative skewness (-0.36) provides evidence a concentration of responses at higher agreement levels, highlighting ethical concerns as a central issue in the adoption of AI in academic environments.

In contrast, the perceived accuracy of GenAI is relatively low ( $M = 1.99$ ,  $SD = 1.06$ ), with a median of 2, indicating limited trust in AI-generated information. The pronounced positive skewness (0.89) suggests that most respondents selected lower values, with fewer observations reflecting high confidence in AI accuracy. This finding indicates a notable scepticism regarding the reliability of GenAI outputs.

The sample composition indicates that 9.26% of respondents identified as Ukrainian, while the remaining 90.74% identified as belonging to other national backgrounds. This distribution provides an opportunity to explore potential differences in perceptions of generative AI across distinct socio-institutional and geopolitical contexts.

The distributional properties across variables, characterised by moderate skewness and negative kurtosis, suggest that responses are relatively dispersed and do not deviate substantially from normality. However, Likert-scale characteristics should be considered when selecting subsequent statistical techniques.

### **RQ1 – To what extent do university staff perceive GenAI as a threat to academic teaching roles?**

The distribution of responses indicates that university staff predominantly do not perceive generative AI as a significant threat to academic teaching roles. A substantial proportion of respondents disagreed with the statement that GenAI will replace university teachers, with 23.55% selecting “strongly disagree” and 27.47% selecting “disagree,” totalling over half of the sample (51.02%). An additional 31.55% adopted a neutral stance, suggesting considerable uncertainty or ambivalence about the future role of AI in academia. In contrast, only a minority of respondents agreed: 12.87% indicated agreement and 4.55% strong agreement. These findings demonstrate that scepticism towards the replacement of university teachers by GenAI clearly outweighs supportive views, with perceptions largely concentrated in disagreement and neutrality rather than endorsement of AI substitution.

### **RQ2 – Do perceptions of GenAI complexity and cost influence beliefs about its ability to replace teachers?**

The results of the Spearman rank correlation analysis indicate that perceptions of GenAI implementation barriers are not significantly associated with beliefs about its potential to replace university teachers. Specifically, the correlation between perceived cost and replacement beliefs is positive but extremely weak and statistically insignificant ( $\rho = 0.037$ ,  $p = 0.345$ ). Similarly, the relationship between perceived complexity and replacement perceptions is also weak and insignificant ( $\rho = 0.042$ ,  $p = 0.288$ ). These findings suggest that respondents’ expectations regarding whether GenAI may substitute academic teaching roles are largely not significantly associated with their assessments of its financial and technical barriers. In other words, perceptions of cost and complexity do not appear to play a decisive role in shaping beliefs about the transformative impact of GenAI on teaching, suggesting that other factors, such as perceived capabilities or ethical considerations, may instead drive these beliefs.

The results of the multiple linear regression analysis indicate that perceived implementation barriers do not significantly influence beliefs about generative AI’s ability to replace university teachers. Although both perceived cost ( $\beta = 0.077$ ,  $p = 0.111$ ) and perceived complexity ( $\beta = 0.021$ ,  $p = 0.650$ ) exhibit positive coefficients, suggesting a weak tendency for higher perceived barriers to be associated with slightly stronger beliefs in AI-driven replacement, neither effect is statistically significant. The model’s explanatory power is extremely limited, as reflected in the very low  $R^2$  value (0.007), indicating that cost and complexity jointly account for less than 1% of the variance in perceived replacement potential. Moreover, the overall model is not statistically significant ( $F = 2.156$ ,  $p = 0.117$ ), further confirming that these factors do not meaningfully shape respondents’ perceptions of whether GenAI may substitute academic teaching roles.

These findings are consistent with the ordinal logistic regression model, which accounts for the ordinal nature of the dependent variable. The estimated coefficients for cost ( $\beta = 0.096$ ,  $p = 0.233$ ) and complexity ( $\beta = 0.033$ ,  $p = 0.662$ ) remain statistically insignificant, confirming the absence of a systematic relationship between perceived implementation barriers and beliefs about AI-driven replacement. The consistency of results across both modelling approaches

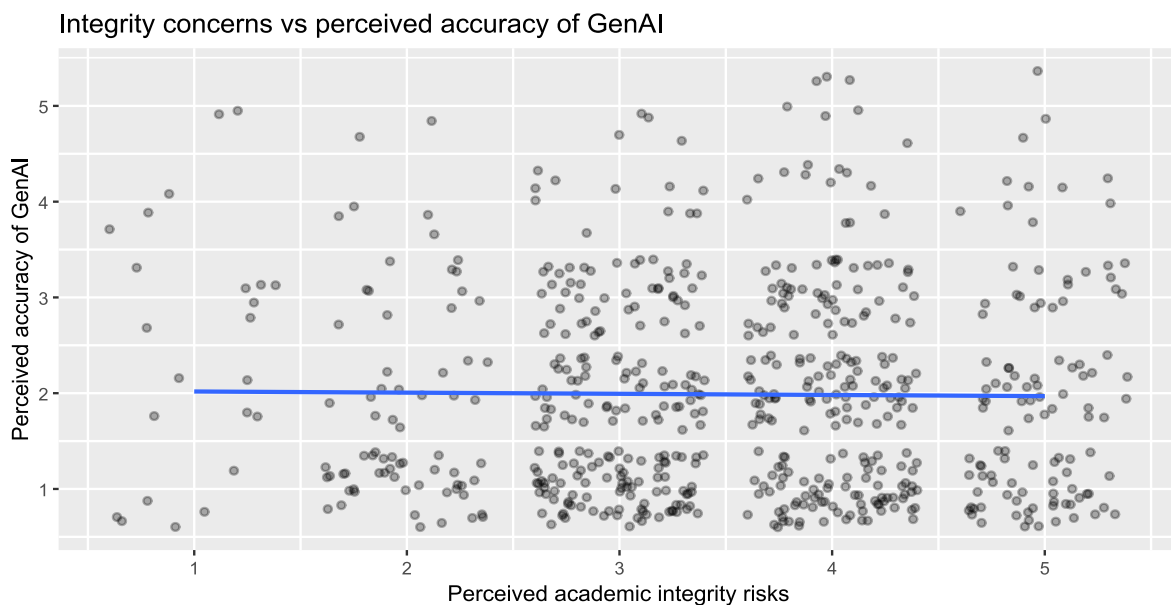
strengthens the robustness of this conclusion. Substantively, this suggests that respondents' evaluations are not significantly associated with perceived financial and technical feasibility. In other words, beliefs about technological substitution appear to be weakly related to practical implementation considerations. They are likely driven instead by other dimensions, such as perceived capabilities, trust in AI outputs, or ethical concerns related to its use in academic environments.

This may indicate a weak conceptual relationship between perceived feasibility constraints and expectations regarding technological substitution.

### **RQ3 – Is the perception of academic integrity violations associated with distrust in GenAI accuracy?**

The correlation analysis provides consistent evidence that perceptions of academic integrity violations are not associated with distrust in the accuracy of generative AI. The Spearman rank correlation reveals a near-zero relationship between perceived integrity risks and perceived accuracy ( $\rho = 0.015$ ,  $p = 0.701$ ), indicating no meaningful monotonic association between the variables. This suggests that respondents' concerns about academic misconduct linked to GenAI do not translate into lower trust in the correctness of AI-generated outputs.

Figure 1 illustrates the relationship between perceived academic integrity risks and perceived accuracy of generative AI. The scatterplot, combined with a fitted linear trend line, demonstrates a near-flat relationship, indicating the absence of a systematic association between the two variables. The wide dispersion of observations across all levels of integrity perception further suggests that respondents who express great concerns about academic misconduct are not consistently more sceptical about AI accuracy. This visual evidence supports the statistical findings, reinforcing the conclusion that ethical concerns and technological trust represent distinct and weakly connected dimensions of GenAI perception.



**Figure 1.** Integrity concerns vs perceived accuracy of GenAI. [Source: Created by the authors]

The findings are further supported by the linear regression analysis (Table 2), which shows that perceived integrity risks have a negligible and statistically insignificant effect on perceived accuracy ( $\beta = -0.012$ ,  $p = 0.772$ ). The model explains only a very small proportion of the variance in accuracy perceptions ( $R^2 \approx 0.0001$ ), confirming that integrity concerns do not account for differences in trust in GenAI. This result indicates that respondents conceptually separate ethical risks from technological performance, rather than interpreting them as directly linked.

**Table 2.** Regression Results: Perceived Accuracy of GenAI as a Function of Academic Integrity Concerns.  
[Source: Created by the authors.]

| Variable            | OLS (Model 1)   | OLS (Model 2)   | Ordinal Logit (Model 3) |
|---------------------|-----------------|-----------------|-------------------------|
| Intercept           | 2.030***(0.152) | 1.638***(0.189) | –                       |
| Integrity           | –0.012(0.041)   | –0.031(0.041)   | –0.002(0.073)           |
| Cost                | –               | 0.030(0.046)    | –                       |
| Complexity          | –               | 0.134**(0.043)  | –                       |
| Observations        | 637             | 637             | 637                     |
| R <sup>2</sup>      | 0.0001          | 0.023           | –                       |
| Adj. R <sup>2</sup> | –0.001          | 0.018           | –                       |
| F-statistic         | 0.084           | 4.957**         | –                       |
| AIC                 | –               | –               | 1685.875                |

*Note.* Dependent variable – Accuracy. Standard errors in parentheses. Significance codes: \*\*\*  $p < 0.001$ , \*\*  $p < 0.01$ , \*  $p < 0.05$ . Model 1 includes integrity only. Model 2 adds cost and complexity. Model 3 reports ordinal logistic regression estimates. OLS estimates are reported for interpretability, while ordinal models are treated as the primary specification for inference.

The robustness of this conclusion is reinforced by the ordinal logistic regression model, which accounts for the ordinal nature of the dependent variable. The effect of integrity remains statistically insignificant ( $\beta = -0.002$ ,  $p = 0.976$ ), providing robust evidence that the absence of association is not driven by model specification. Thus, across both linear and ordinal approaches, no empirical support is found for the hypothesis that concerns about academic integrity are associated with distrust in AI accuracy.

Interestingly, when additional controls are introduced, the extended regression model reveals that perceived complexity has a positive and statistically significant effect on perceived accuracy ( $\beta = 0.134$ ,  $p = 0.002$ ), while integrity and cost remain insignificant. Although the model's explanatory power remains low ( $R^2 \approx 0.023$ ), this suggests that respondents who perceive GenAI as more technically complex may also attribute higher accuracy to it, potentially reflecting a belief that more advanced or sophisticated systems are more reliable. The findings suggest a relative separation between ethical concerns and technological trust, implying that attitudes towards GenAI are multidimensional and not driven by a single evaluative framework.

#### **RQ4 – Do perceptions of GenAI differ across countries or between Ukrainian and non-Ukrainian respondents?**

The sample is clearly unbalanced in terms of respondent background, with 578 non-Ukrainian respondents and 59 Ukrainian respondents. This corresponds to approximately 90.7% and 9.3% of the sample, respectively, which should be considered when interpreting the comparative

results. Despite the smaller Ukrainian subsample, the descriptive statistics reveal several noteworthy differences in the perception of generative AI.

At the descriptive level (Table 3), Ukrainian respondents reported somewhat stronger agreement that GenAI may replace university teachers in the future ( $M = 2.73$ ) than non-Ukrainian respondents ( $M = 2.45$ ). They also perceived the implementation of GenAI as slightly more expensive ( $M = 3.05$  vs  $2.78$ ) and slightly more complex ( $M = 2.86$  vs  $2.79$ ). In addition, Ukrainian respondents expressed somewhat higher trust in AI accuracy ( $M = 2.24$ ) than non-Ukrainian respondents ( $M = 1.96$ ). However, the only dimension on which the two groups differed significantly was academic integrity. Non-Ukrainian respondents reported significantly stronger concern that GenAI leads to violations of academic integrity ( $M = 3.55$ ) than Ukrainian respondents ( $M = 3.14$ ), and this difference was statistically significant at the 5% level ( $t = 2.509$ ,  $p = 0.015$ ). The confidence interval for the mean difference does not include zero, which strengthens the conclusion that this gap is unlikely to be due to random variation.

**Table 3.** Mean values of GenAI perceptions by Ukrainian status. [Source: Created by the authors.]

| Ukrainian status | N   | Replacement | Cost | Complexity | Integrity | Accuracy |
|------------------|-----|-------------|------|------------|-----------|----------|
| No               | 578 | 2.45        | 2.78 | 2.79       | 3.55      | 1.96     |
| Yes              | 59  | 2.73        | 3.05 | 2.86       | 3.14      | 2.24     |

*Note.* All variables are measured on a five-point Likert scale, where 1 = strongly disagree and 5 = strongly agree.

By contrast, the differences observed for replacement, cost, complexity, and accuracy do not reach conventional levels of statistical significance (Table 4). The difference in perceived replacement potential is positive for Ukrainian respondents, but insignificant ( $p = 0.133$ ), suggesting that although Ukrainians appear somewhat more open to the possibility that GenAI could replace teachers, the evidence is not strong enough to conclude that this perception systematically differs from that of non-Ukrainians. A similar pattern holds for perceived cost ( $p = 0.086$ ) and perceived accuracy ( $p = 0.093$ ), which are marginally higher among Ukrainian respondents but remain statistically insignificant at the 5% level. Perceived complexity shows almost no group difference at all ( $p = 0.658$ ). These findings suggest that national background, in terms of Ukrainian versus non-Ukrainian identity, is only weakly associated with perceptions of GenAI, with academic integrity being the only dimension where a robust difference emerges.

**Table 4.** Welch t-tests for differences between Ukrainian and non-Ukrainian respondents. [Source: Created by the authors.]

| Variable    | Mean (No) | Mean (Yes) | t      | df     | p-value | 95% CI of mean difference |
|-------------|-----------|------------|--------|--------|---------|---------------------------|
| Replacement | 2.45      | 2.73       | -1.520 | 65.622 | 0.133   | [-0.649, 0.088]           |
| Cost        | 2.78      | 3.05       | -1.740 | 67.212 | 0.086   | [-0.588, 0.040]           |
| Complexity  | 2.79      | 2.86       | -0.445 | 67.925 | 0.658   | [-0.395, 0.251]           |
| Integrity   | 3.55      | 3.14       | 2.509  | 65.910 | 0.015   | [0.084, 0.741]            |
| Accuracy    | 1.96      | 2.24       | -1.706 | 67.401 | 0.093   | [-0.597, 0.047]           |

*Note.* The only statistically significant difference at the 5% level is observed for academic integrity.

The regression estimates provide evidence of this interpretation. In simple linear models with Ukrainian status as the sole explanatory variable, the coefficient for Ukrainian (Yes) is positive for replacement, cost, complexity, and accuracy, and negative for integrity. This means that Ukrainian respondents tend to report somewhat higher scores on replacement, cost,

complexity, and accuracy, but lower scores on integrity concerns. However, the inferential tests indicate that only the negative coefficient for integrity shows a statistically significant difference. Substantively, this indicates that the Ukrainian/non-Ukrainian distinction does not shape a broad and consistent perception pattern across all dimensions of GenAI, but rather concerns ethical issues related to academic misuse.

### Interpretation of country-level differences

The country-level analysis (Table 5) reveals substantially greater variation than the binary comparison by Ukrainian status. After restricting the sample to countries with at least 10 observations, 19 country categories were retained for analysis. The descriptive statistics demonstrate pronounced heterogeneity across national contexts, indicating that country-specific factors strongly shape perceptions of generative AI. In particular, the perceived potential of GenAI to replace university teachers varies considerably, ranging from very low levels in Canada ( $M = 1.62$ ) to relatively high levels in Latvia ( $M = 3.14$ ). Similarly, perceptions of academic integrity risks exhibit notable cross-country differences, with the lowest values observed in Poland ( $M = 3.03$ ) and the highest in France ( $M = 4.19$ ). The most substantial variation, however, is observed in perceived accuracy, which ranges from very low levels in Canada ( $M = 1.35$ ) and the United States ( $M = 1.45$ ) to considerably higher levels in the Czech Republic ( $M = 2.81$ ) and Latvia ( $M = 2.76$ ). These findings suggest that trust in GenAI outputs differs markedly across national contexts. The results indicate that respondents' evaluations of GenAI are not uniform but embedded in broader institutional, technological, and cultural environments. The observed heterogeneity across all five dimensions, particularly in perceived accuracy and replacement potential, highlights the importance of considering country-specific factors when analysing attitudes towards AI adoption in higher education.

**Table 5.** Descriptive statistics by country (countries with  $N \geq 10$ ). [Source: Created by the authors.]

| Country     | N  | Replacement | Cost | Complexity | Integrity | Accuracy |
|-------------|----|-------------|------|------------|-----------|----------|
| Lithuania   | 46 | 2.41        | 2.91 | 3.17       | 3.76      | 1.85     |
| Romania     | 42 | 2.38        | 3.02 | 2.95       | 3.33      | 2.02     |
| Slovakia    | 41 | 2.56        | 2.90 | 2.93       | 3.07      | 2.39     |
| Canada      | 37 | 1.62        | 3.11 | 3.16       | 3.73      | 1.35     |
| UK          | 35 | 2.14        | 2.77 | 2.57       | 3.54      | 1.77     |
| Other       | 34 | 2.71        | 3.15 | 2.94       | 3.44      | 2.09     |
| Poland      | 34 | 2.85        | 2.71 | 2.53       | 3.03      | 1.97     |
| Germany     | 33 | 2.45        | 2.73 | 2.79       | 3.70      | 2.06     |
| Italy       | 32 | 2.47        | 3.16 | 3.09       | 3.12      | 2.00     |
| France      | 31 | 2.39        | 2.19 | 2.29       | 4.19      | 1.87     |
| Malta       | 30 | 2.73        | 2.90 | 3.13       | 3.90      | 2.63     |
| Switzerland | 30 | 2.43        | 2.77 | 2.73       | 3.60      | 1.53     |
| US          | 29 | 2.24        | 2.34 | 2.59       | 3.76      | 1.45     |
| Bulgaria    | 23 | 2.39        | 2.52 | 2.91       | 3.35      | 2.09     |
| Czech Rep.  | 21 | 2.90        | 3.00 | 2.57       | 3.24      | 2.81     |
| Latvia      | 21 | 3.14        | 3.14 | 2.95       | 3.43      | 2.76     |
| Sweden      | 21 | 2.38        | 2.43 | 2.52       | 3.43      | 1.90     |
| Austria     | 20 | 2.85        | 2.45 | 2.10       | 3.55      | 2.25     |
| Spain       | 13 | 2.69        | 2.54 | 2.69       | 3.85      | 1.92     |

*Note.* All variables are measured on a five-point Likert scale (1 = strongly disagree; 5 = strongly agree).

The descriptive differences are statistically significant according to the Kruskal–Wallis test (Table 6). Significant country-level variation is observed for all five perception variables: replacement ( $\chi^2 = 51.049$ ,  $p < 0.001$ ), cost ( $\chi^2 = 45.763$ ,  $p < 0.001$ ), complexity ( $\chi^2 = 39.608$ ,  $p = 0.002$ ), integrity ( $\chi^2 = 54.557$ ,  $p < 0.001$ ), and accuracy ( $\chi^2 = 70.027$ ,  $p < 0.001$ ). Among these, the strongest cross-country heterogeneity is found for perceived accuracy, followed by integrity and replacement. This means that the way university staff evaluate the benefits and risks of GenAI differs significantly across countries, and these differences are not limited to one isolated perception dimension.

**Table 6.** Kruskal–Wallis tests for cross-country differences. [Source: Created by the authors.]

| Variable    | $\chi^2$ | df | p-value |
|-------------|----------|----|---------|
| Replacement | 51.049   | 18 | <0.001  |
| Cost        | 45.763   | 18 | <0.001  |
| Complexity  | 39.608   | 18 | 0.002   |
| Integrity   | 54.557   | 18 | <0.001  |
| Accuracy    | 70.027   | 18 | <0.001  |

*Note.* Significant country differences are observed for all five GenAI perception variables.

The post hoc Dunn test conducted for the replacement variable refines this conclusion by showing that only a limited number of pairwise differences remain statistically significant after Bonferroni correction. The post hoc analysis (Table 7) reveals that statistically significant differences in replacement perceptions are concentrated around a limited number of country pairs. In particular, Canada consistently differs from several European countries, including Austria, the Czech Republic, Latvia, Malta, Poland, Slovakia, and the aggregated “Other” category. In all significant comparisons, Canadian respondents report substantially lower agreement that GenAI may replace university teachers. In particular, Canada differs significantly from Austria, the Czech Republic, Latvia, Malta, Poland, Slovakia and other countries. Since Canada has one of the lowest mean replacement scores ( $M = 1.62$ ), these results suggest that Canadian respondents are especially sceptical of the idea that GenAI could replace university teachers. In contrast, respondents in some European contexts are relatively more open to this possibility. At the same time, many pairwise contrasts lose significance after correction, indicating that the overall country effect is driven by moderate differences across many countries rather than by widespread sharp contrasts between specific country pairs. The full set of pairwise comparisons (171 tests) is reported in Appendix B Table B1.

**Table 7.** Post hoc Dunn test results for cross-country differences in replacement perceptions. [Source: Created by the authors.]

| Comparison          | Z      | Adjusted p-value |
|---------------------|--------|------------------|
| Austria – Canada    | 4.304  | 0.0029           |
| Canada – Czech Rep. | -4.548 | 0.0009           |
| Canada – Latvia     | -5.378 | <0.001           |
| Canada – Malta      | -4.145 | 0.0058           |
| Canada – Other      | -4.351 | 0.0023           |
| Canada – Poland     | -4.771 | 0.0003           |
| Canada – Slovakia   | -3.962 | 0.0127           |

*Note.* Only statistically significant pairwise differences after Bonferroni correction are reported. All other comparisons are not statistically significant at conventional levels ( $p \geq 0.05$ ).

Importantly, no widespread pattern of significant differences across most country pairs is observed after Bonferroni correction, indicating that the overall country effect is driven primarily by a small number of distinct national contexts rather than systematic divergence across all countries. This suggests that while cross-country variation exists, it is not uniformly distributed but instead reflects specific national outliers.

The linear model with country dummies (Table 8) supports the same substantive reading. Using Austria as the reference category, most coefficients for countries are negative, indicating lower perceived replacement potential in many countries relative to Austria, whereas Latvia shows a positive coefficient. The largest negative coefficients are observed for Canada, the United Kingdom, and the United States, again pointing to comparatively lower agreement in these countries that GenAI may replace academic teachers. Taken together, the results suggest that national context matters much more than Ukrainian identity alone. The cross-country differences are broad, statistically significant, and substantively meaningful, implying that country-specific educational, technological, and cultural environments shape perceptions of GenAI (Palffy Zsuzsanna, 2024).

**Table 8.** Linear model coefficients for replacement by Ukrainian status and country. [Source: Created by the authors.]

| Variable                                                                | Coefficient |
|-------------------------------------------------------------------------|-------------|
| Ukrainian status as a predictor of replacement                          |             |
| Intercept (Non-Ukrainian)                                               | 2.448       |
| Ukrainian                                                               | 0.281       |
| Country effects on replacement perception (reference category: Austria) |             |
| Intercept (Austria)                                                     | 2.850       |
| Bulgaria                                                                | -0.459      |
| Canada                                                                  | -1.228      |
| Czech Rep.                                                              | 0.055       |
| France                                                                  | -0.463      |
| Germany                                                                 | -0.395      |
| Italy                                                                   | -0.381      |
| Latvia                                                                  | 0.293       |
| Lithuania                                                               | -0.437      |
| Malta                                                                   | -0.117      |
| Other                                                                   | -0.144      |
| Poland                                                                  | 0.003       |
| Romania                                                                 | -0.469      |
| Slovakia                                                                | -0.289      |
| Spain                                                                   | -0.158      |
| Sweden                                                                  | -0.469      |
| Switzerland                                                             | -0.417      |
| UK                                                                      | -0.707      |
| US                                                                      | -0.609      |

*Note.* These coefficients are descriptive linear contrasts and should be interpreted relative to Austria as the reference category; no statistical significance testing is applied to these coefficients.

### **RQ5 – Is the perception of generative AI in higher education organised as a unified attitudinal construct, or does it reflect a multidimensional and cognitively differentiated evaluation framework?**

The correlation analysis (Table 9) reveals a weak, only partially coherent structure of perceptions regarding generative AI, indicating limited internal consistency between the



optimistic and sceptical dimensions. The moderate relationship is observed between perceived replacement potential and perceived accuracy ( $\rho = 0.34$ ,  $p < 0.001$ ), suggesting that respondents who believe GenAI is more accurate are also more likely to consider it capable of replacing university teachers. This provides partial support for the existence of an internally consistent “optimistic” dimension centred on perceived technological capability.

**Table 9.** Spearman correlation matrix of GenAI perceptions. [Source: Created by the authors.]

| Variable       | 1      | 2      | 3      | 4      | 5    |
|----------------|--------|--------|--------|--------|------|
| 1. Replacement | 1.00   | 0.34** | 0.04   | 0.04   | 0.08 |
| 2. Accuracy    | 0.34** | 1.00   | 0.07   | 0.14** | 0.02 |
| 3. Cost        | 0.04   | 0.07   | 1.00   | 0.45** | 0.04 |
| 4. Complexity  | 0.04   | 0.14** | 0.45** | 1.00   | 0.10 |
| 5. Integrity   | 0.08   | 0.02   | 0.04   | 0.10   | 1.00 |

*Note.* Spearman rank correlation coefficients are reported. Significance codes: \*\*  $p < 0.01$  (adjusted for multiple comparisons). All variables are measured on a five-point Likert scale.

In contrast, the sceptical dimension exhibits only moderate internal coherence. A statistically significant, relatively strong positive correlation is found between perceived cost and complexity ( $\rho = 0.45$ ,  $p < 0.001$ ), indicating that respondents tend to perceive financial and technical barriers as interrelated. However, the integrity variable shows only weak or negligible associations with the other sceptical indicators (e.g.,  $\rho = 0.10$  with complexity and  $\rho = 0.04$  with cost), suggesting that ethical concerns are not fully integrated into the broader perception of implementation barriers.

Importantly, the relationship between optimistic and sceptical dimensions is largely absent. Most cross-group correlations are weak and statistically insignificant, such as those between replacement and cost ( $\rho = 0.04$ ) and between accuracy and integrity ( $\rho = 0.02$ ). This indicates that respondents do not perceive technological potential and associated risks as opposing or even systematically related dimensions. Instead, these aspects appear to be given little weight.

The findings do not support a strong, internally consistent attitudinal structure that distinguishes between optimistic and sceptical perceptions of GenAI. Rather, respondents' views appear fragmented, with only partial clustering of conceptually related items. The correlation structure further suggests that there is no single, unified attitude towards generative AI among university staff. Rather than evaluating GenAI with a single general positive or negative orientation, respondents appear to distinguish between several relatively weakly associated dimensions. In particular, perceptions related to technological capability, reflected in accuracy and replacement potential, form one interpretative domain, while perceptions of implementation barriers, captured by cost and complexity, form another. Ethical concerns, represented by academic integrity, appear even more detached, showing only weak connections with both capability- and barrier-related assessments. This pattern indicates that attitudes towards GenAI in higher education are cognitively differentiated: respondents do not treat AI as a uniformly beneficial or harmful phenomenon, but instead assess its performance, feasibility, and ethical implications as separate aspects of the same technology.

The reliability analysis (Table 10) reveals weak internal consistency across both the optimistic and sceptical dimensions of GenAI perceptions. The two-item “optimistic” scale, consisting of perceived replacement potential and perceived accuracy, yields a Cronbach's alpha of 0.50, indicating low reliability. Although the correlation between the two items is

moderate ( $r = 0.33$ ), the overall scale does not meet conventional thresholds for acceptable internal consistency. This suggests that respondents do not strongly associate the perceived capability of GenAI with its accuracy, reflecting only partial coherence within the optimistic dimension.

**Table 10.** Reliability analysis of GenAI perception scales. [Source: Created by the authors.]

| Scale                         | Items                       | Cronbach's $\alpha$ | Average $r$ | Interpretation       |
|-------------------------------|-----------------------------|---------------------|-------------|----------------------|
| Optimistic (Capability)       | Replacement, Accuracy       | 0.50                | 0.33        | Weak consistency     |
| Sceptical (Barriers & Risks)  | Cost, Complexity, Integrity | 0.45                | 0.21        | Weak consistency     |
| Sceptical (without Integrity) | Cost, Complexity            | 0.63                | 0.46        | Moderate consistency |

The sceptical dimension, comprising cost, complexity, and integrity, demonstrates even weaker internal consistency, with a Cronbach's alpha of 0.45. While cost and complexity show relatively strong item-total correlations ( $r \approx 0.56$ – $0.60$ ), the integrity variable exhibits a notably low correlation with the overall scale ( $r = 0.15$ ), indicating that ethical concerns are not well integrated with perceptions of implementation barriers. Importantly, the reliability would increase substantially ( $\alpha = 0.63$ ) if integrity were removed, further confirming that this variable represents a distinct perceptual dimension rather than part of a unified sceptical construct.

The factor analysis (Table 11) provides additional evidence of a fragmented attitudinal structure. The Kaiser–Meyer–Olkin (KMO) measure is relatively low ( $MSA = 0.52$ ), indicating weak sampling adequacy and limiting factor reliability for factor analysis. At the same time, Bartlett's test is highly significant ( $\chi^2 = 250.94$ ,  $p < 0.001$ ), confirming that correlations are sufficient for dimensional analysis. However, the parallel analysis suggests that no strong and stable factor structure is clearly supported by the data, further reinforcing the absence of a well-defined latent structure.

**Table 11.** Factor loadings (Two-Factor Solution, Oblimin Rotation). [Source: Created by the authors.]

| Variable    | Factor 1 (Barriers) | Factor 2 (Capability) |
|-------------|---------------------|-----------------------|
| Replacement | –                   | 0.998                 |
| Accuracy    | –                   | 0.324                 |
| Cost        | 0.455               | –                     |
| Complexity  | 0.998               | –                     |
| Integrity   | –                   | –                     |

*Model statistic.*: Variance explained = 46.8%. Factor 1 = 24.7%, Factor 2 = 22.1%

Despite this, a forced two-factor solution reveals a weak and uneven pattern of loadings. The first factor is primarily defined by complexity (loading  $\approx 0.998$ ) and moderately by cost (loading  $\approx 0.455$ ), representing a technical/implementation barrier dimension. The second factor is strongly defined by replacement (loading  $\approx 0.998$ ) and weakly by accuracy (loading  $\approx 0.324$ ), representing a perceived capability dimension. Notably, the integrity variable does not load significantly on either factor, indicating that ethical concerns constitute a distinct, largely independent dimension rather than part of a general scepticism construct.

Figure 2 presents the correlation heatmap of the five perception variables. The visual pattern highlights two relatively stronger clusters: one between replacement and accuracy, and another between cost and complexity. However, the overall correlation structure remains weak, with most relationships close to zero. In particular, the integrity variable shows minimal association with the other variables, reinforcing its conceptual independence. The absence of

strong and consistent correlation patterns supports the conclusion that perceptions of GenAI are fragmented and not governed by a single underlying dimension.

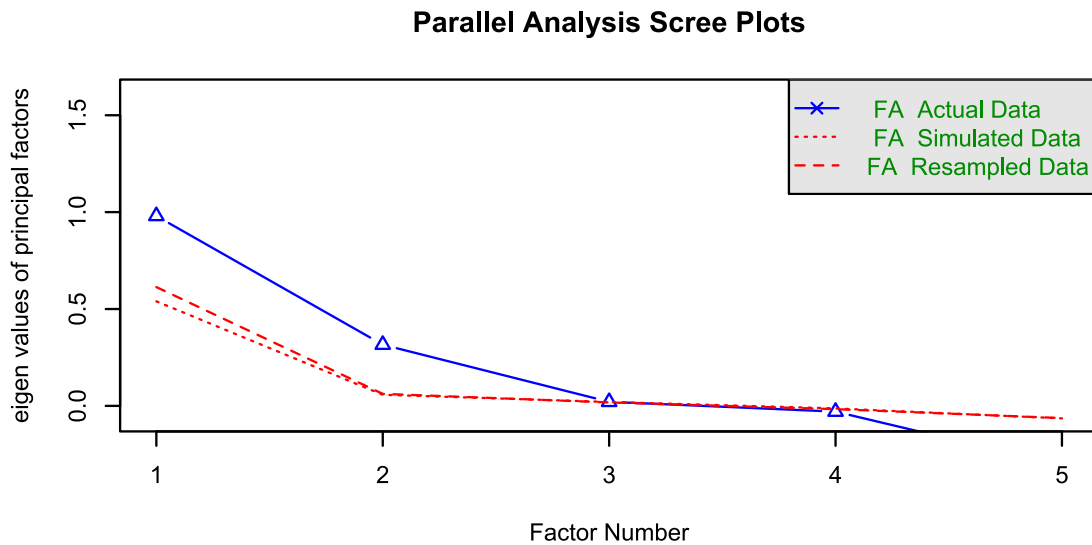


**Figure 2.** Correlation heatmap of GenAI perceptions. [Source: Created by the authors]

Figure 3 illustrates the relationship between perceived academic integrity risks and perceived accuracy of GenAI. The dispersion of observations and the nearly flat regression line indicate the absence of any meaningful relationship between these variables. Respondents who express strong concerns about academic integrity are not systematically more sceptical about AI accuracy. This visual evidence further supports the statistical findings and suggests that ethical concerns constitute a distinct, largely nonsignificant dimension of GenAI perception.

These findings indicate that perceptions of GenAI are not organised into a coherent attitudinal framework. Instead, respondents evaluate capability, implementation barriers, and ethical concerns as largely separate dimensions. This supports the interpretation of a fragmented, multidimensional perception structure rather than a unified, optimistic–sceptical continuum. This interpretation is supported consistently across correlation, reliability, and factor analyses.

The empirical results provide differentiated evidence across the five research questions, confirming that perceptions of generative AI among university staff are neither uniform nor governed by a single evaluative logic. With respect to RQ1, the findings clearly indicate that GenAI is not widely perceived as a direct threat to academic teaching roles. The distribution of responses is strongly concentrated in disagreement and neutrality, with a mean of 2.47 and more than half of respondents rejecting the idea of full teacher replacement. Thus, RQ1 is not supported by a strong perceived substitution effect, but rather points to a cautious or sceptical stance towards the transformative potential of GenAI in teaching.



**Figure 3.** Scatterplot of integrity concerns and perceived accuracy. [Source: Created by the authors]

For RQ2, the results consistently show that perceived implementation barriers, namely cost and complexity, do not significantly influence beliefs about AI-driven replacement. Both correlation and regression analyses yield weak, statistically insignificant relationships, with an explanatory power of close to zero ( $R^2 = 0.007$ ). Therefore, RQ2 is not supported, suggesting that respondents do not base their expectations about AI substitution on practical feasibility considerations. In contrast, RQ3 is also not supported, as no meaningful association is found between perceived academic integrity risks and distrust in AI accuracy ( $\rho = 0.015$ ;  $p = 0.701$ ). This indicates that ethical concerns and technological trust are not significantly associated dimensions, but rather conceptually distinct perceptions.

The findings for RQ4 provide only partial support. While the binary comparison between Ukrainian and non-Ukrainian respondents reveals limited differences, significant only for academic integrity, substantial and statistically significant variation emerges at the country level across all five variables. This indicates that national context plays a much more important role than individual background characteristics in shaping perceptions of GenAI. This is consistent with broader research showing that socio-cultural factors significantly determine how younger cohorts, including Generation Z, form their attitudes towards quality, technology, and innovation-based practices (Spsychalski, 2023), with implications for how institutions design AI governance frameworks. Finally, RQ5 is clearly supported in its alternative formulation: perceptions of GenAI do not form a unified attitudinal construct. Instead, the weak correlations, low reliability coefficients ( $\alpha = 0.50$  and  $0.45$ ), and fragmented factor structure provide evidence that respondents evaluate capability, barriers, and ethical risks as distinct and only loosely connected dimensions. Overall, the results demonstrate that attitudes towards generative AI in higher education are multidimensional, context-dependent, and cognitively differentiated rather than internally consistent.

## DISCUSSION

The findings of this study provide consistent evidence that university staff do not perceive generative artificial intelligence as an immediate substitute for academic teaching roles, but rather as a complementary or uncertain technological development. The low level of perceived replacement potential ( $M = 2.47$ ), combined with a strong concentration of responses in disagreement and neutrality, aligns with prior research emphasising the augmentation rather than the substitution of educators in AI-enhanced learning environments. This supports the argument that AI reshapes pedagogical roles towards facilitation, supervision, and critical engagement rather than full automation (Celik et al., 2022; Jeon & Lee, 2023; Zhang & Fan, 2026). Analogously, in service-oriented organizational contexts, sustainable personnel management emphasizes the renegotiation of roles rather than displacement, reinforcing the view that technological integration demands adaptive human competencies (Vovk & Vovk, 2024). At the same time, the substantial proportion of neutral responses suggests that uncertainty remains a defining feature of staff perceptions, reflecting the ongoing and incomplete institutionalisation of GenAI in higher education (González-Campos et al., 2024; Kohnke & Zou, 2025).

A key contribution of this study is demonstrating that beliefs about AI-driven substitution are not significantly associated with perceived implementation barriers. The absence of a significant relationship between cost, complexity, and replacement perceptions ( $R^2 = 0.007$ ;  $p = 0.117$ ) suggests a relative separation between feasibility considerations and expectations of technological impact. This finding extends prior literature on AI adoption, which typically assumes that perceived barriers play a central role in shaping attitudes towards technology use (Ejdys et al., 2025; Roba & Moulay, 2024). Instead, the results suggest that university staff evaluate the transformative potential of GenAI primarily through cognitive and normative lenses, such as perceived capability or role compatibility, rather than through practical constraints. This reinforces the human-centred perspective that technology adoption in education is driven more by interpretation and meaning-making than by purely technical or economic considerations (Koibichuk et al., 2023; Moravec et al., 2024).

The findings also reveal a pronounced separation between ethical concerns and technological trust. Although academic integrity risks are the most strongly perceived dimension ( $M = 3.51$ ), no significant relationship is found between integrity concerns and perceived accuracy ( $\rho = 0.015$ ;  $p = 0.701$ ). This suggests that ethical evaluations and cognitive assessments operate as distinct dimensions, rather than as mutually reinforcing perceptions. Such a pattern is consistent with the literature on the sociotechnical nature of AI-related risks, in which concerns about misuse, authorship, and responsibility are not necessarily linked to perceptions of system performance (Cotton et al., 2023; Eke, 2023; Yusuf et al., 2024). Moreover, the independence of ethical concerns supports the argument that academic integrity in the context of AI represents a governance and institutional challenge rather than a simple by-product of technological reliability (Artyukhov et al., 2024; Mura & Stehlíková, 2025; Mujtaba, 2025).

Finally, the results provide strong empirical support for the multidimensional and fragmented nature of perceptions of GenAI. Weak correlations, low reliability coefficients ( $\alpha = 0.50$ ;  $\alpha = 0.45$ ), and the absence of a clear factor structure indicate that university staff do not evaluate AI through a single attitudinal framework but instead distinguish between capability, implementation barriers, and ethical risks. This finding directly addresses a gap in the literature, which has often treated attitudes towards AI as internally consistent or unidimensional

constructs (Grabowska & Rzemieniak, 2025; Nazaretsky et al., 2022). In addition, the substantial cross-country heterogeneity observed across all perception dimensions highlights the importance of institutional and contextual factors in shaping AI evaluation, supporting the view that AI adoption is embedded in broader organisational and socio-cultural environments (Kuzior et al., 2024; Welch, 2025; Yarovenko et al., 2024). Overall, the study demonstrates that university staff perceptions of GenAI are cognitively differentiated, context-dependent, and structured around multiple independent evaluative logics rather than a unified orientation towards the technology.

## CONCLUSIONS

This study aimed to examine how university staff cognitively and ethically evaluate GenAI in higher education, with particular attention to perceptions of its potential to replace teaching roles, its practical feasibility, its implications for academic integrity, and its perceived reliability across different national contexts. By focusing on multiple dimensions simultaneously, the study sought to determine whether attitudes towards GenAI form a coherent evaluative framework or instead reflect a fragmented and multidimensional perception structure.

The empirical analysis was based on an anonymous cross-sectional survey of 637 university staff conducted between May and September 2025. The study employed descriptive statistics, Spearman correlations, regression models, non-parametric tests, reliability analysis, and exploratory factor analysis. This multi-method quantitative analytical strategy enabled the examination of five research questions: perceived replacement potential, implementation barriers, ethical concerns, cross-group differences, and the internal consistency of attitudes towards GenAI.

The results reveal several key findings. First, the perceived likelihood of GenAI replacing university teachers is relatively low ( $M = 2.47$ ), with over 51% of respondents disagreeing, indicating limited support for the full automation of teaching roles. Second, academic integrity concerns emerge as the most pronounced dimension ( $M = 3.51$ ), highlighting ethical risks as a central issue in GenAI adoption. Third, trust in GenAI accuracy is comparatively low ( $M = 1.99$ ), suggesting widespread scepticism regarding the reliability of AI-generated outputs. Fourth, perceptions of cost and complexity do not significantly influence beliefs about teacher replacement ( $R^2 = 0.007$ ;  $p = 0.117$ ), suggesting a weak relationship between feasibility and perceived transformative potential. Fifth, the correlation structure demonstrates moderate internal consistency (e.g.,  $\rho = 0.34$  between replacement and accuracy; Cronbach's  $\alpha = 0.50$  and  $0.45$ ), supported by a factor analysis that explains only 46.8% of the variance, confirming that attitudes towards GenAI are fragmented rather than unified.

From a policy perspective, the findings suggest that strategies for integrating generative AI into higher education should move beyond treating attitudes as uniformly positive or negative and instead address distinct perceptual dimensions separately. First, given the strong concern about academic integrity, universities should prioritise developing clear ethical guidelines, redesigning assessments, and implementing AI-aware academic policies that reduce misuse while supporting responsible engagement. Second, the low level of trust in AI accuracy indicates a need for institutional training programmes that improve AI literacy, critical evaluation skills, and understanding of model limitations among academic staff. Third,

because beliefs about AI's transformative role are not driven by perceived cost or complexity, policy frameworks should explicitly address pedagogical implications, rather than focusing solely on technical or financial barriers. Fourth, the significant cross-country heterogeneity suggests that AI governance in higher education should be context-sensitive, considering national institutional environments and digital maturity. Finally, the absence of a unified attitudinal structure implies that effective AI strategies must be multidimensional, simultaneously addressing capability perceptions, implementation constraints, and ethical concerns, rather than relying on a single policy narrative.

## REFERENCES

- Artyukhov, A., Wołowiec, T., Artyukhova, N., Bogacki, S., & Vasylieva, T. (2024). SDG 4, academic integrity and artificial intelligence: Clash or win-win cooperation? *Sustainability*, 16(19), 8483. <https://doi.org/10.3390/su16198483>
- Bilan, S., Šuleř, P., Skrynnyk, O., Krajňáková, E., & Vasilyeva, T. (2022). Systematic bibliometric review of artificial intelligence technology in organizational management, development, change and culture. *Business: Theory and Practice*, 23(1), 1–13. <https://doi.org/10.3846/btp.2022.13204>
- Bîrcă, A., Sandu, C. B., Gogu, E., & Matveiciuc, I. (2025). The assessment of employees' competences on the labour market in the context of the development of artificial intelligence systems. *Montenegrin Journal of Economics*, 21(3), 45–56. <https://doi.org/10.14254/1800-5845/2025.21-3.4>
- Borissov, D., & Liuta, O. (2025). Ethical practices in modern academia: Does length of educational experience and quality of governance contribute to a deeper understanding of academic integrity? *Business Ethics and Leadership*, 9(1), 95–108. [https://doi.org/10.61093/bel.9\(1\).95-108.2025](https://doi.org/10.61093/bel.9(1).95-108.2025)
- Celik, I., Dindar, M., Muukkonen, H., & Järvelä, S. (2022). The promises and challenges of artificial intelligence for teachers: A systematic review of research. *TechTrends*. <https://doi.org/10.1007/s11528-022-00715-y>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2023). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 1–12. <https://doi.org/10.1080/14703297.2023.2190148>
- Derakhshan, A., & Ghiasvand, F. (2024). Is ChatGPT an evil or an angel in second-language education and research? A phenomenographic study of research-active EFL teachers' perceptions. *International Journal of Applied Linguistics*. <https://doi.org/10.1111/ijal.12561>
- Ejdys, J., Garwolińska, M., Lăzăroiu, G., Nica, E., di Pietro, F., Poskrobko, T., & Szpilko, D. (2025). Should we be wary of using artificial intelligence-based big data management in social research? *Journal of Business Economics and Management*, 26(5), 1071–1089. <https://doi.org/10.3846/jbem.2025.24792>
- Eke, D. O. (2023). ChatGPT and the rise of generative AI: Threat to academic integrity? *Journal of Responsible Technology*, 13, 100060. <https://doi.org/10.1016/j.jrt.2023.100060>
- European Commission. (2026). *Guidelines on the ethical use of artificial intelligence and data in teaching and learning for educators* (Rev. ed.). Publications Office of the European Union. <https://doi.org/10.2766/7967834>
- González-Campos, J., López-Núñez, J., & Araya-Pérez, C. (2024). Educación superior e inteligencia artificial: desafíos para la universidad del siglo XXI. *Aloma: Revista de Psicología, Ciències de l'Educació i de l'Esport*, 42(1), 79–90. <https://doi.org/10.51698/aloma.2024.42.1.79-90>
- Grabowska, M., & Rzemieniak, M. (2025). Trust and artificial intelligence – a review of research trends. *Polish Journal of Management Studies*, 32(2), 78–94. <https://doi.org/10.17512/pjms.2025.32.2.05>
- Hutchinson, J. (2025). Navigating AI transformation in healthcare call centers: Balancing efficiency, ethics, and human expertise in HealthConnect's transition. *Health Economics and Management Review*, 6(4), 28–48. <https://doi.org/10.61093/hem.2025.4-03>

- Hwang, G.-J., Xie, H., Wah, B. W., & Gašević, D. (2020). Vision, challenges, roles and research issues of artificial intelligence in education. *Computers and Education: Artificial Intelligence*, 1, 100001. <https://doi.org/10.1016/j.caeai.2020.100001>
- Ioan-Franc, V., & Diamescu, A. M. (2026). The impact of artificial intelligence on humanistic economics. *Amfiteatru Economic*, 28(71), 400–411. <https://doi.org/10.24818/EA/2026/71/400>
- Iordache, R. M., Cioca, V. R., Mihaila, D., Štreimikienė, D., & Ionescu, Ș. E. (2025). An analysis on leadership and decision making errors in the new artificial intelligence influenced organizational environment. *Polish Journal of Management Studies*, 31(2), 123–137. <https://doi.org/10.17512/pjms.2025.31.2.08>
- Jeon, J., & Lee, S. (2023). Large language models in education: A focus on the complementary relationship between human teachers and ChatGPT. *Education and Information Technologies*. <https://doi.org/10.1007/s10639-023-11834-1>
- Kaczorowska-Spychalska, D., Kotula, N., Mazurek, G., & Sułkowski, Ł. (2024). Generative AI as source of change of knowledge management paradigm. *Human Technology*, 20(1), 131–154. <https://doi.org/10.14254/1795-6889.2024.20-1.7>
- Kohnke, L., & Zou, D. (2025). Artificial intelligence integration in TESOL teacher education: Promoting a critical lens guided by TPACK and SAMR. *TESOL Quarterly*. <https://doi.org/10.1002/tesq.3396>
- Koibichuk, V., Samoilikova, A., & Vasylieva, T. (2023). Digitalization and innovation transfer as a leadership trend in education: Bibliometric analysis and social analytics. *Springer Proceedings in Business and Economics*, 233–247. [https://doi.org/10.1007/978-3-031-28131-0\\_17](https://doi.org/10.1007/978-3-031-28131-0_17)
- Kozhakhmetova, A., Mamyrbayev, A., Zhidebekkyzy, A., & Bilan, S. (2024). Assessing the impact of artificial intelligence on project efficiency enhancement. *Knowledge and Performance Management*, 8(2), 109–126. [https://doi.org/10.21511/kpm.08\(2\).2024.09](https://doi.org/10.21511/kpm.08(2).2024.09)
- Kundu, A., Bej, T., & Mondal, R. (2025). Exploring teachers' AI literacy: Cognitive, pedagogical, ethical, and contextual insights from Indian schools. *Teaching Education*, 1–21. <https://doi.org/10.1080/10476210.2025.2570227>
- Kuzior, A., Sira, M., Zozuľáková, V., & Hetenyi, M. (2024). Navigating AI regulation: A comparative analysis of EU and US legal frameworks. *Materials Research Proceedings*, 45, 258–266. <https://doi.org/10.21741/9781644903315-30>
- Mong, D. D., & Thanh, H. P. (2024). Relationship between artificial intelligence and legal education: A bibliometric analysis. *Knowledge and Performance Management*, 8(2), 13–27. [https://doi.org/10.21511/kpm.08\(2\).2024.02](https://doi.org/10.21511/kpm.08(2).2024.02)
- Moravec, V., Hynek, N., Gavurova, B., & Kubak, M. (2024). Everyday artificial intelligence unveiled: Societal awareness of technological transformation. *Oeconomia Copernicana*, 15(2), 367–406. <https://doi.org/10.24136/oc.2961>
- Mura, L., & Stehlíková, B. (2025). The ethics of artificial intelligence: Safeguarding human dignity, social justice and environmental stability in the age of AI. *Equilibrium. Quarterly Journal of Economics and Economic Policy*, 20(2), 479–507. <https://doi.org/10.24136/eq.3743>
- Mujtaba, B. G. (2025). Human-AI intersection: Understanding the ethical challenges, opportunities, and governance protocols for a changing data-driven digital world. *Business Ethics and Leadership*, 9(1), 109–126. [https://doi.org/10.61093/bel.9\(1\).109-126.2025](https://doi.org/10.61093/bel.9(1).109-126.2025)
- Nazaretsky, T., Ariely, M., Cukurova, M., & Alexandron, G. (2022). Teachers' trust in AI-powered educational technology and a professional development program to improve it. *British Journal of Educational Technology*. <https://doi.org/10.1111/bjet.13232>
- Némethová, I., Ablonczy-Mihályka, L., Kecskés, P., & Stradiotová, E. (2025). Leadership readiness in higher education: A SILDM-based competency diagnosis. *European Journal of Interdisciplinary Studies*, 17(2), 44–60. <https://doi.org/10.24818/ejis.2025.13>
- Nguyen, A., Ngo, H. N., Hong, Y., Dang, B., & Nguyen, B.-P. T. (2022). Ethical principles for artificial intelligence in education. *Education and Information Technologies*. <https://doi.org/10.1007/s10639-022-11316-w>



- OECD. (n.d.). AI principles. <https://www.oecd.org/en/topics/ai-principles.html>
- Pálffy, Zs., & Mihályka, L. A. (2024). Cultural influence on sustainability discourse: A comparative analysis of CEO letters in the cosmetics and fashion industry. *Journal of Ecohumanism*, 3(4), 2779–2792. <https://doi.org/10.62754/joe.v3i4.3796>
- Pálffy, Zs. (2024). Companies local embeddedness and culture: A review. *The Economic Research Guardian*, 14(2).
- Pálffy, Zs., & Ablonczy-Mihalyka, L. (2024). Varying levels of trust across multi-level governance: A sustainability perspective. *Chemical Engineering Transactions*, 114, 1045–1050. <https://doi.org/10.3303/CET24114175>
- Pálffy, Zs., Ablonczy-Mihalyka, L., & Kecskés, P. (2023). A sustainable model of corporate embeddedness: Navigating through a fuzzy concept. *Chemical Engineering Transactions*, 107, 163–168. <https://doi.org/10.3303/CET23107028>
- Pisică, A. I., & Zaharia, R. M. (2025). What artificial intelligence means to us – students' insights. *European Journal of Interdisciplinary Studies*, 17(1), 142–152. <https://doi.org/10.24818/ejis.2025.09>
- Popa, I., Cioc, M. M., Breazu, A., & Popa, C. F. (2024). Identifying sufficient and necessary competencies in the effective use of artificial intelligence technologies. *Amfiteatru Economic*, 26(65), 33–52. <https://doi.org/10.24818/EA/2024/65/33>
- Qian, Z. (2021). Applications, risks and countermeasures of artificial intelligence in education. In *2021 2nd International Conference on Artificial Intelligence and Education (ICAIE)*. IEEE. <https://doi.org/10.1109/icaie53562.2021.00026>
- Rahmanov, F., Ganiyeva, S., Aliyeva, N., Neymatova, L., & Aghazada, T. (2025). The impact of education digitalization on achieving SDG4: A comparative assessment of Azerbaijan and SDG4 leaders. *Problems and Perspectives in Management*, 23(2), 634–650. [https://doi.org/10.21511/ppm.23\(2\).2025.46](https://doi.org/10.21511/ppm.23(2).2025.46)
- Rajasekaran, S., Adam, T., & Tilmes, K. (2024). *Digital pathways for education: Enabling greater impact for all*. World Bank. <http://hdl.handle.net/10986/42386>
- Roba, M., & Moulay, O. K. (2024). Risk management in using artificial neural networks. *SocioEconomic Challenges*, 8(2), 302–313. [https://doi.org/10.61093/sec.8\(2\).302-313.2024](https://doi.org/10.61093/sec.8(2).302-313.2024)
- Schinello, S. (2025). Challenges and opportunities in the use of artificial intelligence in creative economy: Insights from expert interviews. *Economics and Sociology*, 18(1), 199–216. <https://doi.org/10.14254/2071-789X.2025/18-1/10>
- Škare, M., Porada-Rochon, M., & Veselica Celić, R. (2025). Total factor productivity dynamics and the artificial intelligence paradox: Evidence from long-memory analysis. *Journal of International Studies*, 18(4), 218–242. <https://doi.org/10.14254/2071-8330.2025/18-4/11>
- Skydan, M. (2024). Strategic directions of the development of higher education in Ukraine. *Social and Labour Relations: Theory and Practice*, 14(1), 1–11. [https://doi.org/10.21511/slntp.14\(1\).2024.01](https://doi.org/10.21511/slntp.14(1).2024.01)
- Spychalski, B. (2023). Socio-cultural factors shaping the attitude of Generation Z and Generation Alpha youth towards quality. *Journal of Sustainable Development of Transport and Logistics*, 8(2), 298–311. <https://doi.org/10.14254/jsdtl.2023.8-2.23>
- Sullivan, M., Kelly, A., & McLaughlan, P. (2023). ChatGPT in higher education: Considerations for academic integrity and student learning. *Journal of Applied Learning & Teaching*, 6(1). <https://doi.org/10.37074/jalt.2023.6.1.17>
- United Nations Educational, Scientific and Cultural Organization (UNESCO). (2022). *Recommendation on the ethics of artificial intelligence* (SHS/BIO/PI/2021/1). <https://unesdoc.unesco.org/ark:/48223/pf0000381137>
- United Nations Educational, Scientific and Cultural Organization (UNESCO). (2023). *Guidance for generative AI in education and research*. <https://doi.org/10.54675/EWZM9535>
- Vovk, I., & Vovk, Y. (2024). Sustainable personnel management in the hospitality industry: Enhancing organizational performance through employee engagement and commitment. *Economics, Management and Sustainability*, 9(2), 44–58. <https://doi.org/10.14254/jems.2024.9-2.4>

- Welch, N. (2025). Exploring underrepresented employee perceptions of AI receptivity through leaders' stakeholder engagement. *Business Ethics and Leadership*, 9(2), 108–119. [https://doi.org/10.61093/bel.9\(2\).108-119.2025](https://doi.org/10.61093/bel.9(2).108-119.2025)
- Yarovenko, H., Kuzior, A., Norek, T., & Lopatka, A. (2024). The future of artificial intelligence: Fear, hope or indifference? *Human Technology*, 20(3), 611–639. <https://doi.org/10.14254/1795-6889.2024.20-3.10>
- Yusuf, A., Pervin, N., & Román-González, M. (2024). Generative AI and the future of higher education: A threat to academic integrity or reformation? Evidence from multicultural perspectives. *International Journal of Educational Technology in Higher Education*, 21(1). <https://doi.org/10.1186/s41239-024-00453-6>
- Zhang, Q., & Fan, J. (2026). Beyond substitution: How teacher interpersonal behaviours empower students to communicate with AI – A dual-mediation model of enjoyment and trust. *European Journal of Education*, 61(1). <https://doi.org/10.1111/ejed.70499>
- Zimosz, P., & Ober, J. (2025). Impact of artificial intelligence on education. *Knowledge Economy and Lifelong Learning*, 1(2), 77–98. [https://doi.org/10.61093/kell.1\(2\).77-98.2025](https://doi.org/10.61093/kell.1(2).77-98.2025)
- Zuo, Z., Luo, Y., Yan, S., & Jiang, L. (2025). From perception to practice: Artificial intelligence as a pathway to enhancing digital literacy in higher education teaching. *Systems*, 13(8), 664. <https://doi.org/10.3390/systems13080664>

---

## Authors' Note

All correspondence should be addressed to  
 Serhiy Lyeonov  
 Silesian University of Technology, Poland; Lithuania Business College, Lithuania  
 Email address: [Serhiy.Lyeonov@polsl.pl](mailto:Serhiy.Lyeonov@polsl.pl)

---

*Human Technology*  
 ISSN 1795-6889  
<https://ht.csr-pub.eu>

## **Appendix**

### **Appendix A. Questionnaire**

#### **Survey instrument used in the final analysis**

##### *Instructions to respondents*

Please indicate the extent of your agreement with each statement.

Scale: 1 = strongly disagree, 2 = disagree, 3 = neither agree nor disagree, 4 = agree, 5 = strongly agree

##### *Background/context items*

1. Country where your organisation is located
2. I am Ukrainian (Yes/No)

##### *Perception items*

3. GenAI will replace university teachers in the future.
4. The implementation of GenAI in universities is expensive.
5. Using GenAI is complex and requires significant technical knowledge.
6. The use of GenAI leads to violations of academic integrity.
7. GenAI knows everything and always provides accurate information.

#### **Questionnaire note for ethics and anonymity**

The questionnaire was administered anonymously. It did not collect names, email addresses, phone numbers, institutional identifiers, addresses, IP-linked respondent identifiers, or any other direct personal identifiers. It also did not include questions on health, religion, ethnicity, political preferences, income, sexual orientation, or other sensitive categories of personal data. The collected information was limited to professional perceptions of GenAI, country of organisational location, and a non-identifying binary contextual item on whether the respondent was Ukrainian.

The survey was open from 12 May to 10 September 2025.

## Appendix B

Table B1. Full Dunn Test Structure

| Comparison             | Z        | P.unadj  | P.adj    |
|------------------------|----------|----------|----------|
| Austria - Bulgaria     | 1,42138  | 0,155206 | 1        |
| Austria - Canada       | 4,303818 | 1,68E-05 | 0,002871 |
| Bulgaria - Canada      | 2,861822 | 0,004212 | 0,720274 |
| Austria - Czech Rep.   | -0,15391 | 0,877681 | 1        |
| Bulgaria - Czech Rep.  | -1,59916 | 0,109785 | 1        |
| Canada - Czech Rep.    | -4,54792 | 5,42E-06 | 0,000926 |
| Austria - France       | 1,42907  | 0,152984 | 1        |
| Bulgaria - France      | -0,08978 | 0,928458 | 1        |
| Canada - France        | -3,22239 | 0,001271 | 0,217388 |
| Czech Rep. - France    | 1,620358 | 0,105155 | 1        |
| Austria - Germany      | 1,499454 | 0,133756 | 1        |
| Bulgaria - Germany     | -0,03558 | 0,97162  | 1        |
| Canada - Germany       | -3,21403 | 0,001309 | 0,223812 |
| Czech Rep. - Germany   | 1,694458 | 0,090178 | 1        |
| France - Germany       | 0,060152 | 0,952034 | 1        |
| Austria - Italy        | 1,398457 | 0,161976 | 1        |
| Bulgaria - Italy       | -0,13152 | 0,895361 | 1        |
| Canada - Italy         | -3,29672 | 0,000978 | 0,167274 |
| Czech Rep. - Italy     | 1,590639 | 0,111691 | 1        |
| France - Italy         | -0,04462 | 0,964409 | 1        |
| Germany - Italy        | -0,10597 | 0,915608 | 1        |
| Austria - Latvia       | -0,87962 | 0,379064 | 1        |
| Bulgaria - Latvia      | -2,3504  | 0,018753 | 1        |
| Canada - Latvia        | -5,37783 | 7,54E-08 | 1,29E-05 |
| Czech Rep. - Latvia    | -0,73473 | 0,462506 | 1        |
| France - Latvia        | -2,42263 | 0,015409 | 1        |
| Germany - Latvia       | -2,50673 | 0,012185 | 1        |
| Italy - Latvia         | -2,39802 | 0,016484 | 1        |
| Austria - Lithuania    | 1,566684 | 0,117189 | 1        |
| Bulgaria - Lithuania   | -0,05855 | 0,95331  | 1        |
| Canada - Lithuania     | -3,50879 | 0,00045  | 0,076977 |
| Czech Rep. - Lithuania | 1,775941 | 0,075743 | 1        |
| France - Lithuania     | 0,041986 | 0,96651  | 1        |
| Germany - Lithuania    | -0,02318 | 0,981503 | 1        |
| Italy - Lithuania      | 0,091234 | 0,927306 | 1        |
| Latvia - Lithuania     | 2,6369   | 0,008367 | 1        |
| Austria - Malta        | 0,610108 | 0,54179  | 1        |
| Bulgaria - Malta       | -0,93254 | 0,351057 | 1        |
| Canada - Malta         | -4,14496 | 3,40E-05 | 0,005812 |
| Czech Rep. - Malta     | 0,788028 | 0,43068  | 1        |
| France - Malta         | -0,91268 | 0,361413 | 1        |
| Germany - Malta        | -0,98623 | 0,32402  | 1        |
| Italy - Malta          | -0,87552 | 0,38129  | 1        |
| Latvia - Malta         | 1,584951 | 0,112978 | 1        |
| Lithuania - Malta      | -1,0376  | 0,299454 | 1        |
| Austria - Other        | 0,570388 | 0,568414 | 1        |
| Bulgaria - Other       | -1,01429 | 0,310444 | 1        |
| Canada - Other         | -4,35131 | 1,35E-05 | 0,002314 |
| Czech Rep. - Other     | 0,752398 | 0,451812 | 1        |
| France - Other         | -1,00321 | 0,315761 | 1        |
| Germany - Other        | -1,08107 | 0,279667 | 1        |
| Italy - Other          | -0,96585 | 0,334118 | 1        |
| Latvia - Other         | 1,569354 | 0,116565 | 1        |
| Lithuania - Other      | -1,14468 | 0,252342 | 1        |
| Malta - Other          | -0,06143 | 0,951019 | 1        |
| Austria - Poland       | 0,216996 | 0,828212 | 1        |

|                       |          |          |          |
|-----------------------|----------|----------|----------|
| Bulgaria - Poland     | -1,38315 | 0,166618 | 1        |
| Canada - Poland       | -4,7705  | 1,84E-06 | 0,000314 |
| Czech Rep. - Poland   | 0,393585 | 0,693887 | 1        |
| France - Poland       | -1,40423 | 0,160252 | 1        |
| Germany - Poland      | -1,4886  | 0,136594 | 1        |
| Italy - Poland        | -1,37019 | 0,170628 | 1        |
| Latvia - Poland       | 1,210542 | 0,226071 | 1        |
| Lithuania - Poland    | -1,585   | 0,112965 | 1        |
| Malta - Poland        | -0,45899 | 0,646238 | 1        |
| Other - Poland        | -0,41061 | 0,681362 | 1        |
| Austria - Romania     | 1,56041  | 0,118663 | 1        |
| Bulgaria - Romania    | -0,04104 | 0,967267 | 1        |
| Canada - Romania      | -3,41749 | 0,000632 | 0,108075 |
| Czech Rep. - Romania  | 1,766132 | 0,077374 | 1        |
| France - Romania      | 0,059396 | 0,952637 | 1        |
| Germany - Romania     | -0,00422 | 0,996634 | 1        |
| Italy - Romania       | 0,107861 | 0,914106 | 1        |
| Latvia - Romania      | 2,614521 | 0,008935 | 1        |
| Lithuania - Romania   | 0,020184 | 0,983896 | 1        |
| Malta - Romania       | 1,036655 | 0,299897 | 1        |
| Other - Romania       | 1,140866 | 0,253926 | 1        |
| Poland - Romania      | 1,572542 | 0,115825 | 1        |
| Austria - Slovakia    | 1,085336 | 0,277773 | 1        |
| Bulgaria - Slovakia   | -0,53185 | 0,594831 | 1        |
| Canada - Slovakia     | -3,96223 | 7,43E-05 | 0,012697 |
| Czech Rep. - Slovakia | 1,282334 | 0,199725 | 1        |
| France - Slovakia     | -0,47833 | 0,632418 | 1        |
| Germany - Slovakia    | -0,55113 | 0,581542 | 1        |
| Italy - Slovakia      | -0,43497 | 0,663587 | 1        |
| Latvia - Slovakia     | 2,127296 | 0,033396 | 1        |
| Lithuania - Slovakia  | -0,57549 | 0,56496  | 1        |
| Malta - Slovakia      | 0,499039 | 0,617752 | 1        |
| Other - Slovakia      | 0,583244 | 0,559729 | 1        |
| Poland - Slovakia     | 1,012583 | 0,311259 | 1        |
| Romania - Slovakia    | -0,58262 | 0,560152 | 1        |
| Austria - Spain       | 0,380054 | 0,703905 | 1        |
| Bulgaria - Spain      | -0,86221 | 0,388573 | 1        |
| Canada - Spain        | -3,28483 | 0,00102  | 0,174496 |
| Czech Rep. - Spain    | 0,519932 | 0,603111 | 1        |
| France - Spain        | -0,83065 | 0,406172 | 1        |
| Germany - Spain       | -0,88413 | 0,376624 | 1        |
| Italy - Spain         | -0,80032 | 0,423526 | 1        |
| Latvia - Spain        | 1,162432 | 0,24506  | 1        |
| Lithuania - Spain     | -0,90487 | 0,365535 | 1        |
| Malta - Spain         | -0,12264 | 0,902389 | 1        |
| Other - Spain         | -0,0777  | 0,938067 | 1        |
| Poland - Spain        | 0,227696 | 0,819882 | 1        |
| Romania - Spain       | -0,90909 | 0,363301 | 1        |
| Slovakia - Spain      | -0,50463 | 0,61382  | 1        |
| Austria - Sweden      | 1,320385 | 0,186707 | 1        |
| Bulgaria - Sweden     | -0,073   | 0,941804 | 1        |
| Canada - Sweden       | -2,86196 | 0,00421  | 0,719959 |
| Czech Rep. - Sweden   | 1,49261  | 0,135539 | 1        |
| France - Sweden       | 0,009465 | 0,992448 | 1        |
| Germany - Sweden      | -0,04432 | 0,964653 | 1        |
| Italy - Sweden        | 0,049567 | 0,960468 | 1        |
| Latvia - Sweden       | 2,227337 | 0,025925 | 1        |
| Lithuania - Sweden    | -0,02689 | 0,978548 | 1        |
| Malta - Sweden        | 0,830935 | 0,40601  | 1        |
| Other - Sweden        | 0,907263 | 0,364268 | 1        |
| Poland - Sweden       | 1,266076 | 0,205486 | 1        |

|                          |          |          |          |
|--------------------------|----------|----------|----------|
| Romania - Sweden         | -0,04261 | 0,966009 | 1        |
| Slovakia - Sweden        | 0,43422  | 0,664129 | 1        |
| Spain - Sweden           | 0,785317 | 0,432268 | 1        |
| Austria - Switzerland    | 1,388447 | 0,165001 | 1        |
| Bulgaria - Switzerland   | -0,12183 | 0,903032 | 1        |
| Canada - Switzerland     | -3,23042 | 0,001236 | 0,211371 |
| Czech Rep. - Switzerland | 1,57773  | 0,114628 | 1        |
| France - Switzerland     | -0,03536 | 0,971791 | 1        |
| Germany - Switzerland    | -0,09554 | 0,923883 | 1        |
| Italy - Switzerland      | 0,008612 | 0,993129 | 1        |
| Latvia - Switzerland     | 2,374653 | 0,017565 | 1        |
| Lithuania - Switzerland  | -0,08017 | 0,936105 | 1        |
| Malta - Switzerland      | 0,870209 | 0,384186 | 1        |
| Other - Switzerland      | 0,958419 | 0,337852 | 1        |
| Poland - Switzerland     | 1,355986 | 0,175104 | 1        |
| Romania - Switzerland    | -0,09672 | 0,922948 | 1        |
| Slovakia - Switzerland   | 0,436154 | 0,662725 | 1        |
| Spain - Switzerland      | 0,799313 | 0,424109 | 1        |
| Sweden - Switzerland     | -0,04123 | 0,96711  | 1        |
| Austria - UK             | 2,26133  | 0,023739 | 1        |
| Bulgaria - UK            | 0,742448 | 0,457816 | 1        |
| Canada - UK              | -2,37753 | 0,017429 | 1        |
| Czech Rep. - UK          | 2,470605 | 0,013488 | 1        |
| France - UK              | 0,908209 | 0,363768 | 1        |
| Germany - UK             | 0,861159 | 0,389151 | 1        |
| Italy - UK               | 0,961806 | 0,336147 | 1        |
| Latvia - UK              | 3,292055 | 0,000995 | 0,170074 |
| Lithuania - UK           | 0,955153 | 0,3395   | 1        |
| Malta - UK               | 1,839746 | 0,065805 | 1        |
| Other - UK               | 1,964842 | 0,049432 | 1        |
| Poland - UK              | 2,378413 | 0,017387 | 1        |
| Romania - UK             | 0,917264 | 0,359004 | 1        |
| Slovakia - UK            | 1,468028 | 0,142097 | 1        |
| Spain - UK               | 1,534687 | 0,124861 | 1        |
| Sweden - UK              | 0,801817 | 0,422659 | 1        |
| Switzerland - UK         | 0,936687 | 0,348919 | 1        |
| Austria - US             | 1,898318 | 0,057654 | 1        |
| Bulgaria - US            | 0,419704 | 0,674701 | 1        |
| Canada - US              | -2,59144 | 0,009558 | 1        |
| Czech Rep. - US          | 2,093472 | 0,036307 | 1        |
| France - US              | 0,549258 | 0,582828 | 1        |
| Germany - US             | 0,498374 | 0,618221 | 1        |
| Italy - US               | 0,597314 | 0,550298 | 1        |
| Latvia - US              | 2,884797 | 0,003917 | 0,66975  |
| Lithuania - US           | 0,557292 | 0,577328 | 1        |
| Malta - US               | 1,442467 | 0,149171 | 1        |
| Other - US               | 1,546947 | 0,121876 | 1        |
| Poland - US              | 1,940922 | 0,052268 | 1        |
| Romania - US             | 0,529463 | 0,596484 | 1        |
| Slovakia - US            | 1,054011 | 0,291878 | 1        |
| Spain - US               | 1,24744  | 0,212236 | 1        |
| Sweden - US              | 0,485882 | 0,627051 | 1        |
| Switzerland - US         | 0,579664 | 0,562141 | 1        |
| UK - US                  | -0,32696 | 0,743701 | 1        |